

Riscos cibernéticos do Uso de Inteligência Artificial Generativa na Elaboração de Despachos e Decisões no Sistema Judiciário Brasileiro

Luciana Muniz Costa, Marcus Aurélio Carvalho Georg, Rafael Rabelo Nunes

lmunizc@gmail.com; macgeorg@gmail.com; rafaelrabelo@unb.br

1 Discente na Universidade de Brasília (UnB), Campus Darcy Ribeiro, Departamento de Engenharia Elétrica, CEP 70910-900 Brasília-DF Brasil

2 Superior Tribunal de Justiça

3 Docente na Universidade de Brasília (UnB), Campus Darcy Ribeiro, Departamento de Engenharia Elétrica, CEP 70910-900 Brasília-DF Brasil

Resumo: *O Poder Judiciário brasileiro passou por uma transformação significativa com a incorporação da tecnologia, especialmente com o processo eletrônico, que aumentou a velocidade e a produtividade dos julgamentos. A pandemia acelerou essa evolução, introduzindo videoconferências, audiências virtuais, citações eletrônicas, atendimento remoto e o uso de inteligência artificial na análise processual. Tudo isso potencializou riscos de vazamento de informações, divulgação antecipada de decisões e interrupções na prestação jurisdicional, entre outros. A inserção da inteligência artificial generativa pode intensificar esses riscos. No contexto jurídico, o uso da inteligência artificial generativa em tarefas como redação de despachos, análise de casos e previsão de decisões exige configuração e monitoramento rigorosos. Este estudo qualitativo e exploratório correlaciona riscos de inteligência artificial generativa com os principais riscos de negócio. Para isso foi utilizada a técnica bow-tie para correlacionar os principais riscos de negócio com os riscos relacionados com a adoção da inteligência artificial generativa. A partir da relação entre riscos de negócio e riscos cibernéticos associados ao uso da inteligência artificial generativa, o estudo propõe diretrizes para seleção de controles prioritários na adoção dessa tecnologia no processo de elaboração de despachos e de decisões no sistema judiciário brasileiro.*

Palavras-chave: *inteligência artificial generativa; judiciário; riscos cibernéticos; riscos de inteligência artificial generativa.*

Cyber risks of the use of generative artificial intelligence in the elaboration of dispatches and decisions in the Brazilian judiciary

Abstract: *The Brazilian Judiciary has undergone a significant transformation with the incorporation of technology, especially with the electronic process, which increased the speed and productivity of judgments. The pandemic accelerated this evolution, introducing video conferences, virtual hearings, electronic citations, remote care and the use of artificial intelligence in procedural analysis. All of this potentiated risks of leakage of information, early dissemination of decisions and interruptions in the judicial provision, among others. The insertion of general artificial intelligence can intensify these risks. In the legal context, the use of general artificial intelligence in tasks such as order writing, case analysis and decision forecasting requires strict configuration and monitoring. This qualitative and exploratory study correlates risks of general artificial intelligence with the main business risks. For this, the bow-Tie technique was used to correlate the main business risks with the risks related to the adoption of generative artificial intelligence. From the relation between business risks and cyber risks associated with the use of generative artificial intelligence, the study proposes guidelines for the selection of priority controls in the adoption of this technology in the process elaboration of dispatches and decisions in the Brazilian judiciary.*

Keywords: *Generative Artificial Intelligence; Judiciary. Cyber risk; Generative Artificial Intelligence.*

1. Introdução

Impulsionada pela adoção de tecnologias inovadoras, a Justiça brasileira tem passado por transformações significativas, com foco em eficiência, agilidade e acessibilidade. A informatização dos processos judiciais, eliminou o uso de papel e acelerou a tramitação, alterou profundamente a rotina do Judiciário e da advocacia (Hino & Cunha, 2020).

Neste cenário, o Brasil evidencia um compromisso progressivo com a incorporação de tecnologias emergentes, em especial a inteligência artificial (IA), por meio de políticas e programas orientados à modernização do sistema de justiça (STF, 2023). Nesse contexto, o Conselho Nacional de Justiça (CNJ) assume posição estratégica ao normatizar e fomentar a aplicação de soluções tecnológicas no âmbito do Poder Judiciário, com vistas a consolidar maior eficiência, transparência e legitimidade na prestação jurisdicional.

A adoção da IA, inclusive da Inteligência Artificial Generativa, traz benefícios como automação de tarefas e aumento da produtividade, mas também amplia riscos já existentes. Entre eles, destacam-se a segurança cibernética, a proteção de dados e os vieses algorítmicos, que podem comprometer a imparcialidade das decisões judiciais (Alves, Georg & Nunes, 2023). Nesse sentido, destaca-se que o *National Institute of Standards and Technology* (NIST), publicou doze riscos da disseminação de uso de IA Generativa. Alguns dos riscos são confabulação, recomendações perigosas ou violentas, privacidade de dados, toxicidade, viés e homogeneização, entre outros (NIST, 2024).

Nesse sentido, este estudo tem como principal objetivo relacionar os principais fatores de risco que podem afetar o processo de despachos e decisões no sistema judiciário, com os riscos no uso de IA Generativa propostos com o framework NIST AI-600-1 (NIST, 2024). Trata-se de um estudo de natureza aplicada, qualitativa, com objetivos exploratórios. Como procedimento técnico foi utilizada análise documental do *framework* e dos riscos de negócio (Alves, Georg & Nunes, 2023). Os resultados obtidos permitiram relacionar os 12 riscos do uso de IA Generativa, com os 10 riscos de negócio do processo de despachos e decisões no sistema judiciário. Com isso, tornou-se possível propor a seleção de controles para mitigar os referidos riscos de negócio traçando diretrizes para a implantação segura e responsável desse tipo de tecnologia nas plataformas e sistemas judiciais.

Este trabalho está organizado da seguinte forma: a Seção 2 apresenta os principais conceitos relacionados ao trabalho; a Seção 3 explica a metodologia adotada para a condução do estudo; a Seção 4 relaciona os riscos de disseminação de uso de IA Generativa aos principais riscos de negócio do processo de despachos e decisões no sistema judiciário brasileiro.

2. Referencial Teórico

Nessa seção apresentam-se os conceitos necessários para que se compreenda esse trabalho. Primeiramente, discorre-se sobre Inteligência Artificial na Seção 2.1. Na

sequência, a Seção 2.2 trata sobre Governança. A Seção 2.3 aborda a adoção da inteligência artificial no Judiciário Brasileiro e a Seção 2.4 trata da estrutura da Publicação NISTAI-600-1 voltada para a gestão de riscos de IA Generativa. Além disso, a Seção 2.5 apresenta alguns trabalhos correlatos a esta pesquisa.

2.1 Inteligência Artificial (IA)

A IA, ou *Artificial Intelligence* (AI), é um ramo da computação voltado à criação de sistemas capazes de simular comportamentos humanos, como tomada de decisão, aprendizado e resolução de problemas (Russell & Norvig, 2010).

A Associação Brasileira de Normas Técnicas (ABNT, 2023) descreve, na ISO/IEC 22989:2023, a disciplina de IA como pesquisa e desenvolvimento de mecanismos e aplicações de sistemas de IA. A pesquisa e o desenvolvimento podem ocorrer em diversos campos, como ciência da computação, ciência de dados, humanidades, matemática e ciências naturais. Conceitua também o sistema de IA como sistema desenvolvido que gera saídas como conteúdo, previsões, recomendações ou decisões para um determinado conjunto de objetivos definidos pelo homem. O sistema desenvolvido pode utilizar diversas técnicas e abordagens relacionadas à inteligência artificial para desenvolver um modelo para representar dados, conhecimento, processos etc., que podem ser realizados para realizar tarefas (ABNT 2023).

Do ponto de vista filosófico, a viabilidade das máquinas que possuem inteligência tem sido debatida. Esse debate levou à introdução de dois tipos diferentes de IA: a denominada IA fraca e a IA forte. Na IA fraca, o sistema só pode processar símbolos (letras, números etc.) sem nunca entender o que faz. Na IA forte, o sistema também processa símbolos, mas entende verdadeiramente o que faz. As denominações “IA fraca” e “IA forte” são principalmente importantes para os filósofos, mas irrelevantes para pesquisadores e profissionais de IA. Após esse debate, apareceram as qualificações de “IA restrita” versus “IA generalista”, que são mais adequadas ao campo da IA. Um sistema de “IA restrita” é capaz de resolver tarefas definidas para resolver um problema específico. Um sistema de “IA generalista” aborda uma ampla gama de tarefas definidas com um nível satisfatório de desempenho. Os sistemas atuais de IA são considerados “restritos”. Ainda não se sabe se os sistemas de “IA generalistas” serão tecnicamente viáveis no futuro (ABNT, 2023).

Outra norma da ABNT descreve a inteligência artificial como sistemas capazes de aprendizado contínuo, tomada de decisão automatizada e adaptação comportamental, exigindo estratégias específicas de gestão para garantir segurança e responsabilidade (ABNT, 2024).

A denominada “IA Generativa” não é o mesmo que “IA Generalista” (ABNT, 2023). A IA Generativa pode ser compreendida como uma categoria de técnicas de IA que aprendem a partir de dados existentes e geram novos conteúdos, como texto, imagens, áudio ou outras modalidades, que se assemelham aos dados de treinamento.

Diferentemente dos modelos discriminativos, que apenas classificam ou fazem previsões com base em dados de entrada, os modelos generativos criam novas instâncias de dados que compartilham características com os dados originais (NIST, 2023). A IA Generativa é uma tecnologia que tem ganhado espaço por sua acessibilidade e potencial de automação, permitindo desde tarefas simples até o desenvolvimento de modelos personalizados, mesmo por usuários com conhecimento técnico limitado. Embora a IA abranja diversas abordagens, a IA Generativa destaca-se por sua capacidade de criar conteúdo sintético, como textos, imagens, vídeos e áudios, a partir de dados de entrada (NIST, 2024). A Figura 1 destaca que a IA Generativa representa uma parte do espectro da IA (Gartner, 2024)

Figura 1 – IA Generativa no espectro da IA



Gartner, 2024

O uso isolado da IA Generativa pode gerar expectativas excessivas e resultados abaixo do esperado. O sucesso na adoção depende da capacidade das organizações de superar limitações técnicas e operacionais, garantindo que a IA Generativa gere valor real e sustentável. Nesse cenário, o uso responsável e bem orientado da IA-G representa um diferencial competitivo e institucional relevante (Gartner, 2024).

Fora do Brasil, a normatização da IA em nível global tem sido conduzida por diversas organizações internacionais que buscam estabelecer diretrizes para o desenvolvimento dessa tecnologia, como: OCDE, IEEE, União Europeia, Unesco, ISO/IEC e NIST.

2.2 Governança

No contexto atual, os princípios da governança corporativa, conforme definidos pelo Instituto Brasileiro de Governança Corporativa (IBGC), incluem integridade, transparência, equidade, prestação de contas e responsabilidade corporativa.

Esses princípios, que fundamentam a direção ética e eficiente das organizações privadas, encontram paralelo nos princípios da governança pública estabelecidos pelo

Decreto nº 9.203/2017, que destacam a capacidade de resposta, integridade, confiabilidade, melhoria regulatória, prestação de contas e transparência. Essa convergência indica que tanto no âmbito privado quanto no público, a governança pauta-se em valores que asseguram responsabilidade, transparência e ética no uso do poder e na prestação de serviços.

Além disso, esses princípios são aplicáveis para a governança de outros domínios como tecnologia da informação, privacidade, segurança da informação, segurança cibernética, dados e IA, onde a transparência nas decisões automatizadas, a responsabilidade pelos impactos sociais e éticos, e a equidade no tratamento dos envolvidos são essenciais para garantir a confiança, sustentabilidade e conformidade regulatória nas tecnologias disruptivas (IBGC, 2023; Governo Federal, 2017).

Quando se trata de IA, a governança se baseia em três pilares fundamentais que garantem a qualidade, segurança e ética no uso e gestão dos dados: i - o primeiro pilar, governança de dados para IA, concentra-se no controle e na administração dos dados usados no desenvolvimento dos sistemas, assegurando a integridade, privacidade e relevância dos dados coletados; ii - o segundo pilar, governança de dados com IA, utiliza a própria IA para aprimorar os processos de gestão, como automação de classificação e detecção de anomalias, otimizando assim a eficiência e a segurança; e iii - o terceiro pilar, governança de dados de IA, trata da gestão dos dados gerados pelos próprios sistemas de IA, como logs e resultados, garantindo auditabilidade, explicabilidade e responsabilidade. Essa estrutura integrada promove sistemas de IA mais confiáveis e alinhados a princípios éticos essenciais à era digital (Pinto, 2025). No campo da IA Generativa responsável, a governança refere-se à aplicação de valores por meio de regras, práticas e processos que alinhem o uso da tecnologia à mitigação de riscos (CNJ, 2024).

As tecnologias de IA têm um potencial significativo para transformar a sociedade e a vida das pessoas, desde o comércio e a saúde até o transporte, a cibersegurança, o meio ambiente e o próprio planeta. A IA pode impulsionar o crescimento econômico inclusivo e apoiar avanços científicos que melhorem as condições do mundo (NIST, 2024). Os sistemas de IA oferecem benefícios significativos e sua adoção exige que as organizações revisem seus objetivos de governança, uso de dados e valores institucionais, considerando os impactos potenciais dessas tecnologias (ABNT, 2023).

No entanto, também apresenta riscos que podem impactar negativamente indivíduos, grupos, organizações, comunidades, a sociedade, o meio ambiente e o planeta. Assim como em outras tecnologias, os riscos da IA podem surgir de diversas formas e ser caracterizados como de curto ou longo prazo, de alta ou baixa probabilidade, sistêmicos ou localizados, e de alto ou baixo impacto (NIST, 2023).

A IA, por si só, não é inerentemente ética ou antiética, justa ou enviesada, essas características decorrem de seu uso e das decisões humanas envolvidas em seu desenvolvimento, treinamento e aplicação (Gartner, 2024). Esses sistemas dependem de

grandes volumes de dados digitalizados, frequentemente marcados por vieses e lacunas de representatividade, que podem ser reproduzidos ou amplificados nos resultados. Esses vieses, originados por dados ou algoritmos tendenciosos, comprometem a precisão dos sistemas e podem gerar impactos negativos na inclusão social e no desempenho organizacional (Holdsworth, 2023).

Estudos mostram que modelos de IA podem gerar narrativas falsas e conteúdos discriminatórios, afetando diretamente a justiça e a qualidade das decisões judiciais (Bender, Gebru & McMillan-Major, 2021). Além disso, a IA Generativa apresenta o dilema do duplo uso, podendo ser aplicada tanto para fins benéficos quanto maliciosos, como fraudes, manipulação social e disseminação de desinformação (Barrett et al., 2023; Tredinnick & Laybats, 2023; Wach et al., 2023).

Diante disso, a governança e o gerenciamento de riscos tornam-se essenciais para aproveitar os benefícios da IA Generativa e mitigar impactos negativos. Como suas capacidades ainda são menos compreendidas que as de outras formas de IA, a IA Generativa exige níveis diferenciados de supervisão e configurações específicas de interação humano-máquina para garantir uma aplicação ética, segura e eficaz (NIST, 2024).

2.3. Adoção da Inteligência Artificial no Judiciário brasileiro

A evolução da IA Generativa tem impulsionado sua adoção em setores diversos, incluindo o Judiciário, com aplicações que vão desde a geração automatizada de documentos até a análise preditiva de litígios e otimização de fluxos de trabalho.

O relatório do Conselho Nacional de Justiça (CNJ, 2024) evidencia o avanço da digitalização no Judiciário brasileiro, com cerca de 83 milhões de processos eletrônicos (dados de 10/2024). Entre as iniciativas, destaca-se a criação do Centro de Prevenção, Tratamento e Resposta a Incidentes Cibernéticos, parte da Estratégia Nacional de Segurança Cibernética, voltada a fortalecer a proteção de dados e a resiliência institucional. O aumento expressivo de decisões e despachos — até 6,2 milhões e 7,7 milhões por mês, respectivamente — reflete ganhos de produtividade, impulsionados pela adoção de tecnologias digitais e IA, especialmente no âmbito do Programa Justiça 4.0 (CNJ, 2024). O uso de IA e IA-G tem automatizado tarefas repetitivas, agilizado triagens e auxiliado na elaboração de minutas, promovendo eficiência e celeridade.

No entanto, o relatório do CNJ (2024) alerta para riscos éticos e de governança, como opacidade, viés, violação de privacidade e impactos ambientais. A dependência crescente de IA-G no Judiciário exige governança rigorosa, com políticas claras, auditorias frequentes e controles robustos. Os riscos recorrentes — como confabulação, viés algorítmico e exposição de dados — podem comprometer a equidade, a imparcialidade e a segurança jurídica. A literatura e o relatório do CNJ reforçam que, embora a IA traga ganhos substanciais, sua aplicação no sistema de justiça deve ser cuidadosamente

monitorada para garantir decisões justas, seguras e contextualizadas (CNJ, 2024; Alves, Georg & Nunes, 2023).

2.4. Publicação NIST voltada para a gestão de riscos cibernéticos de AI-G

O NIST lançou o *AI Risk Management Framework* (AI RMF) em 2023, que fornece uma estrutura para a gestão de riscos em sistemas de inteligência artificial. O AI RMF considera um sistema de IA como um sistema projetado ou baseado em máquina que, para um determinado conjunto de objetivos, gera saídas como previsões, recomendações ou decisões que influenciam ambientes reais ou virtuais. Os sistemas de IA são concebidos para operar com diferentes níveis de autonomia (Adaptado de: OECD Recommendation on AI:2019; ISO/IEC 22989:2022).

Em 2024, o NIST publicou o *Framework* de Gestão de Riscos de Inteligência Artificial: Perfil de Inteligência Artificial Generativa - NIST.AI.600-1, que constitui um perfil complementar ao AI RMF. O referido perfil foi desenvolvido em resposta à crescente complexidade e impacto dos modelos generativos, e busca orientar organizações na identificação, categorização e mitigação de riscos emergentes. O documento permite adaptação a diferentes contextos institucionais e regulatórios.

A estrutura do perfil é composta por três seções principais: (i) introdução ao escopo e à finalidade do documento; (ii) descrição de doze riscos únicos ou agravados pela IA Generativa; e (iii) conjunto de ações organizadas conforme as funções: *Govern*: Governança (GV), *Map*: Mapeamento (MP), *Measure*: Medição (MS) e *Manage*: Gestão (MG). Cada ação é vinculada a subcategorias e inclui identificadores, riscos associados, palavras-chave e atores de IA relevantes.

A abordagem proposta contribui para o fortalecimento da governança de IA Generativa ao oferecer uma estrutura modular, orientada por funções e alinhada aos princípios de confiabilidade, segurança e transparência. Sua aplicabilidade permite a implementação de controles de riscos envolvidos na adoção de IA Generativa.

2.5.Trabalhos Correlatos

A discussão sobre riscos no Poder Judiciário brasileiro tem ganhado destaque, especialmente com a aceleração da transformação digital. Diversos estudos recentes têm abordado os riscos associados à adoção da IA-G no sistema de justiça. Além dos ataques diretos, a própria arquitetura do ecossistema de IA introduz riscos sistêmicos. A dependência de modelos pré-treinados e bibliotecas de código aberto cria complexas vulnerabilidades na cadeia de suprimentos (*supply chain*), onde ameaças podem ser herdadas de componentes de terceiros sem o conhecimento do usuário final. Soma-se a isso o risco inerente à geração de conteúdo, como as "alucinações" (informações factualmente incorretas) e a criação de deepfakes (Chesney & Citron, 2019).

De forma igualmente alarmante, estudos sobre envenenamento de dados de treinamento (Wallace et al., 2021) revelam a possibilidade de corromper um modelo em sua origem, inserindo "backdoors" que podem ser ativados para sabotar análises jurídicas ou introduzir vieses deliberados. Oliveira (2022) enfatiza os desafios éticos e normativos da regulação da IA ressaltando a necessidade de uma abordagem adaptativa que considere os impactos sociais e jurídicos da tecnologia. Esses trabalhos reforçam a importância de uma governança robusta e de mecanismos de controle para mitigar os riscos da IA-G em decisões judiciais. Hoch e Engelmann (2023) apontam a ausência de um marco regulatório no Brasil como um fator de vulnerabilidade, sugerindo diretrizes inspiradas no modelo europeu para garantir transparência e imparcialidade. Pesquisas sobre injeção de prompt, como as de Greshake et al. (2023), demonstram como instruções maliciosas, ocultas em documentos aparentemente inofensivos, podem coagir um sistema de IA a ignorar suas salvaguardas e exfiltrar dados sigilosos.

O trabalho de Alves, Georg e Nunes (2023) mapeou e classificou os principais riscos de negócio associados às atividades essenciais do Judiciário. O estudo ofereceu uma base para o desenvolvimento de estratégias de mitigação, destacando vulnerabilidades em processos críticos como a elaboração de decisões e despachos. Almada e Zanatta (2024) destacam que, embora a IA Generativa ofereça avanços significativos na automação de tarefas jurídicas e na pesquisa legal, ela também pode amplificar injustiças estruturais e gerar erros graves, especialmente quando utilizada sem supervisão humana adequada.

Não se identificaram estudos que objetivassem relacionar os principais fatores de risco que possam afetar o processo de despachos e decisões no sistema judiciário, com os riscos no uso de IA Generativa, tornando-se possível propor a seleção de controles para mitigar os referidos riscos de negócio traçando diretrizes para a implantação segura e responsável desse tipo de tecnologia nas plataformas e sistemas judiciais.

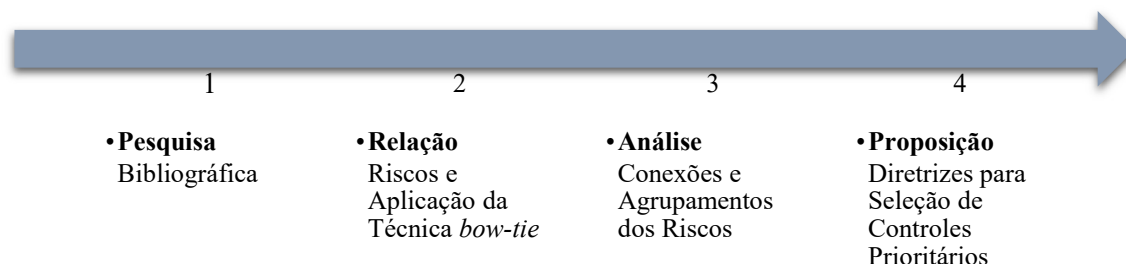
É precisamente diante dessa lacuna que se insere o presente trabalho, pois a ascensão da IA Generativa inaugura uma nova dimensão de ameaças, extrapolando os riscos tradicionais. Essa emergência tecnológica impõe níveis inéditos de urgência e complexidade ao gerenciamento de riscos, agravados pelo cenário em que a multiplicidade de controles e salvaguardas disponíveis contrasta com a limitação de recursos organizacionais para sua efetiva implementação.

3. Metodologia

Este estudo é uma pesquisa de natureza aplicada, pois se concentra na aplicação prática e na solução de problemas específicos. Adotou-se uma abordagem qualitativa para interpretar e compreender um fenômeno do mundo real, sujeito à subjetividade e nem sempre representável em números (Silva & Menezes, 2005, p. 20).

A pesquisa possui objetivo exploratório, visando proporcionar uma melhor compreensão do problema. O trabalho seguiu as etapas descritas na Figura 2.

Figura 2 – Passos deste trabalho



Fonte: Elaborado pelos autores

3.1. Passo 1 - Pesquisa Bibliográfica

Este passo consistiu em um levantamento da literatura sobre os riscos de negócio no contexto do Sistema Judiciário Brasileiro e riscos associados ao uso de IA Generativa. Foram analisados estudos anteriores, como os de Alves, Georg & Nunes (2023) e Alves et al. (2025), que mapearam riscos ligados à produção de decisões judiciais e ao funcionamento de sistemas judiciais informatizados.

Sequencialmente, foram realizados levantamentos e análises na literatura científica, normas técnicas e documentos institucionais relevantes para fundamentar o estudo. As referências contemplaram documentos estruturantes, como a publicação NIST AI-6001 (NIST, 2024), , que identificou os 12 (doze) principais riscos de inteligência artificial generativa, as normas da Associação Brasileira de Normas Técnicas - ABNT, incluindo ABNT NBR ISO/IEC 22989 (2023), que define conceitos e terminologia em inteligência artificial, e ABNT NBR ISO/IEC 42001 (2024), que detalha o sistema de gestão para IA. Foram consideradas também as normas ABNT NBR ISO 27001 e 27005, relacionadas à segurança da informação, e as normas ABNT NBR ISO 31000 e 31010, que tratam respectivamente da gestão e avaliação de riscos, com destaque para a utilização da técnica *bow-tie* na análise de ameaças e barreiras preventivas.

Posteriormente, as diretrizes para a governança pública foram embasadas no Decreto nº 9.203, de 22 de novembro de 2017, que estabelece a política de governança da administração pública federal, definindo princípios essenciais como capacidade de resposta, integridade, confiabilidade, melhoria regulatória, prestação de contas, responsabilidade e transparência (Brasil, 2017). Ademais, os princípios da governança corporativa, conforme definidos pelo Instituto Brasileiro de Governança Corporativa, também foram considerados para alinhar as práticas éticas e responsáveis do estudo (IBGC, 2020).

Ademais, foi aplicado o *heat map* como abordagem metodológica para representar visualmente a densidade de interconexões entre riscos identificados na análise *bow-tie* e os riscos de IA Generativa, permitindo associar a conectividade à criticidade dos

processos e apoiar a proposição da seleção de controles para mitigar os referidos riscos de negócio traçando diretrizes para a implantação segura e responsável desse tipo de tecnologia nas plataformas e sistemas judiciais.

Também foram consideradas publicações relevantes da literatura internacional, como Arrieta et al. (2020), Bender et al. (2021) e Brundage et al. (2018), que abordam riscos emergentes, impactos sociais e limites éticos da IA. Além disso, a pesquisa incorporou contribuições atuais da 'Trilha de Implantação de Governança de Dados para IA (Pinto, 2025), que oferecem uma abordagem prática e integradora para a governança de dados e de IA enfatizando os pilares da governança de IA e o papel da ética, da conformidade e da gestão estruturada de riscos.

Além disso, a pesquisa analisou documentos institucionais, como a Estratégia Brasileira de Inteligência Artificial (EBIA), que orienta a implementação responsável, segura e ética da IA no País, alinhando iniciativas públicas aos princípios de transparência, respeito aos direitos fundamentais e promoção da inovação, bem como, diretrizes e recomendações do Conselho Nacional de Justiça (CNJ, 2024) acerca da adoção e governança de IA nos tribunais brasileiros, compondo assim um referencial amplo e atualizado para a análise dos riscos e definição de controles no âmbito judicial.

3.2. Passo 2 - Relações entre os Riscos e Aplicação da Técnica *Bow-Tie*

Com base no material coletado, realizou-se o relacionamento entre os dez principais riscos de negócio do processo de despachos e decisões no sistema judiciário brasileiro (Alves, Georg & Nunes ,2023), apresentados na Tabela 1, e os doze riscos associados à IA Generativa (NIST,2024), dispostos na Tabela 2.

Tabela 1 – Riscos de Negócio

N	Descrição
1	Divulgação antecipada de votos, determinações ou decisões
2	Vazamento de informações sigilosas, protegidas por segredo de Justiça ou dados pessoais
3	Emissão ou alteração não autorizada de determinações ou decisões
4	Interrupção da prestação jurisdicional
5	Previsibilidade ou manipulação da distribuição dos processos
6	Perda de informações
7	Parcialidade ou favorecimentos pessoais
8	Assuntos indesejados ou inadequados em determinações e decisões
9	Julgamentos legítimos, porém, com base em elementos adulterados

Alves, Georg & Nunes (2023)

Tabela 2 – Riscos de Inteligência Artificial Generativa

N	Descrição
1	Informações CBRN
2	Confabulação
3	Recomendações Perigosas ou Violentas
4	Privacidade de Dados
5	Impacto Ambiental
6	Configuração Humano-IA
7	Integridade da Informação
8	Segurança da Informação
9	Propriedade Intelectual
10	Conteúdo Obsceno, Degradante e/ou Abusivo
11	Toxicidade, Viés e Homogeneização
12	Cadeia de Valor e Integração de Componentes

NIST.AI.600-1 (2024)

A etapa envolveu o estabelecimento da relação entre esses dois conjuntos. Para isso, empregou-se a técnica *bow-tie*, que permite visualizar de forma clara as relações causais entre causas (vulnerabilidades), eventos centrais e consequências (impactos) do processo de despachos e decisões do sistema judiciário, além de identificar barreiras preventivas e mitigadoras (controles e funções NIST, 2024).

3.3. Passo 3 - Análise das Conexões e Agrupamentos dos Riscos

Sequencialmente, foi construído um *heat map* que viabilizou uma análise das conexões entre os riscos identificados. A visualização permitiu a identificação de padrões de recorrência e correlações, evidenciando distintos graus de interdependência entre os riscos de negócio e os riscos associados à IA Generativa, possibilitando um agrupamento dos riscos.

Para isso, este trabalho adotou como premissas: i. a classificação dos grupos de conectividade foi realizada com base em intervalos percentuais fixos, sem aplicação de pesos diferenciados entre os critérios considerados. Dessa forma, cada componente da conectividade contribui de forma igual para o resultado; e ii. o número de interconexões entre os riscos está diretamente relacionado ao impacto no negócio, ou seja, quanto maior o volume de conexões estabelecidas, maior tende a ser a relevância da conectividade.

Essa abordagem permite associar a densidade de conexões à criticidade dos processos envolvidos. Opcionalmente, pode-se adotar a fórmula: $\text{Impacto} = (\text{Interconexões} * w1) + (\text{Consequências} * w2)$, onde: $w1$ e $w2$ são pesos atribuídos conforme a importância relativa de cada fator.

3.4. Passo 4 – Diretrizes para Seleção de Controles Priorizados

A aplicação da técnica *bow-tie* à análise da relação entre os riscos de negócio e riscos de IA Generativa, combinada com os padrões de recorrência revelados pelo *heat map*, possibilitou o agrupamento estruturado dos riscos.

A classificação dos grupos de conectividade foi realizada com base em intervalos percentuais fixos, sem aplicação de pesos diferenciados entre os critérios considerados. Considerando a hipótese da organização dos riscos em 4 (quatro) grupos distintos, foram identificadas combinações, que levaram à composição de cenários possíveis. A análise resultou em diretrizes para seleção de controles prioritários, com base na intensidade das conexões entre os riscos. Essa abordagem permite associar a densidade de conexões à criticidade dos processos envolvidos. Caso necessário, pode-se adotar $\text{Impacto} = (\text{Interconexões} * w_1)$, onde: w_1 é peso atribuído conforme a importância relativa de cada fator.

Dessa forma, cada componente da conectividade contribui de forma igual para o resultado. Para ilustrar a aplicabilidade das diretrizes propostas, a Tabela 3 apresenta um modelo conceitual.

Tabela 3 – Proposição de Diretrizes para Seleção de Controles Prioritários

Seq	Tema	Explicação
1º	Contexto	Estabelecimento do contexto
2º	Riscos de Negócio	Identificação dos principais de negócio como eventos centrais
3º	Relação entre Riscos	Identificação de cada risco de negócio como evento central e aplicação da técnica <i>bow-tie</i> associando os riscos de IA Generativa como causas, riscos de negócio como consequências (impactos) e os respectivos controles, dentro de suas funções, como barreiras preventivas e mitigadoras
4º	Interseções entre Riscos	Identificação/análise de interseções entre os riscos (mapa de calor).
5º	Agrupamento dos Riscos	Classificação dos grupos com base nos intervalos propostos • Grupo 1: $\geq 60\%$ → Alta Conectividade • Grupo 2: 30–59% → Média Conectividade • Grupo 3: 1–29% → Baixa Conectividade • Grupo 4: 0% → Sem Conectividade
6º	Seleção de Controles	Controles associados aos grupos, considerando: Ordenação dos grupos em ordem crescente; riscos de negócio ordenados de forma decrescente de acordo com a relação com participação em mais de um grupo; riscos de negócio ordenados de forma decrescente de acordo com o maior o número de conexões; e riscos de negócio com maior número de riscos de negócio relacionados como consequências.

Fonte: Elaborado pelos autores

Para ilustrar a aplicabilidade das diretrizes propostas, o modelo conceitual foi exemplificado na Tabela 4 no contexto específico do processo de despachos e decisões

judiciais, evidenciando sua relevância prática na identificação e priorização de controles voltados à mitigação dos riscos associados ao uso de IA Generativa.

Tabela 4 – Aplicação das Diretrizes Propostas para Seleção de Controles Prioritários.

Seq	Tema	Aplicação
1º	Contexto	Processo de Despacho e Decisões no Sistema Judiciário Brasileiro Tabela 1 – Riscos de Negócio (RNEG's)
2º	Relação entre Riscos	Tabela 2 – Riscos de IA Generativa (RGAI's) Tabela A1 (Anexo A) - e aplicação da Técnica <i>bow-tie</i>
3º	Conexões entre Riscos	Figura B1 (Anexo B) – Aplicação da Técnica <i>Heat Map</i> Tabela C1 (Anexo C) – Conexões e Grupos
4º	Agrupamento dos Riscos	Tabela D1 (Anexo D) – Características dos Grupos
5º	Cenários de Priorização	Tabela E1 (Anexo E) – Cenários possíveis de combinações para a seleção de controles priorizados
6º	Seleção de Controles	Tabela F1 (Anexo F) – Aplicação da hipótese de adoção do cenário 1 como critério de priorização. Tabela 4 – Riscos e Seleção de Controles Priorizados

Fonte: Elaborado pelos autores

4. Resultados

A análise realizada possibilitou o estabelecimento de correspondências diretas e indiretas entre os dez riscos de negócio relacionados ao processo de despachos e decisões no sistema judiciário brasileiro (Alves, Georg & Nunes, 2023) e os doze riscos associados ao uso de IA Generativa (NIST, 2024). Como resultado, verificou-se a existência de áreas de sobreposição relevantes, especialmente em dimensões associadas à confiabilidade, à segurança da informação, à transparência e à responsabilização. Essa correlação evidenciou não apenas a aderência dos riscos tecnológicos aos já identificados no contexto processual do Poder Judiciário, mas também trouxe à tona lacunas emergentes vinculadas ao uso de modelos de linguagem generativa, como riscos de vieses algorítmicos e de degradação da qualidade decisória. O mapeamento sistemático obtido foi sintetizado nas tabelas comparativas apresentadas na seção seguinte, constituindo insumo essencial para a proposição, em estágio posterior, de controles de mitigação voltados à implantação segura e responsável da IA Generativa em sistemas judiciais.

A principal contribuição deste estudo consiste na proposição de diretrizes para a seleção de controles prioritários voltados à prevenção e mitigação de riscos relacionados à adoção de IA-G no processo de despachos e decisões judiciais no Brasil. Essas diretrizes estão apresentadas na Tabela 1, como proposição, e na Tabela 2, como exemplo de aplicação. Para fundamentar essa proposta, foram empregadas duas abordagens metodológicas complementares: a técnica *bow-tie*, recomendada pela norma ABNT NBR ISO 31010 para avaliação de riscos, e o *heat map*, técnica engenhosa de visualização que

tem sido amplamente empregada em estudos de análise de padrões e agrupamentos em dados complexos (Wilkinson & Friendly, 2009).

. A técnica *bow-tie* permitiu estruturar as relações causais entre vulnerabilidades da IA Generativa (causas) e suas possíveis consequências (impactos), possibilitando a identificação de controles preventivos e mitigadores (NIST, 2024), conforme demonstrado na Tabela A1 (ver Anexo A). Em complemento, o *heat map* representa visualmente a densidade de interconexões entre os riscos, permitindo associar a conectividade à criticidade dos processos e apoiar a priorização de controles, conforme demonstrado na Figura B1 (ver Anexo B).

A combinação das técnicas *bow-tie* e *heat map* mostrou-se relevante para embasar a proposição dos quatro grupos de conectividade, ao permitir uma visualização clara dos riscos associados à IA-G e dos respectivos controles, evidenciando os respectivos cruzamentos com os riscos de negócio. Enquanto a técnica *bow-tie* possibilitou estruturar as relações causais entre vulnerabilidades e consequências, o *heat map* contribuiu para quantificar e destacar visualmente a densidade de interconexões entre os riscos, facilitando a análise da criticidade dos processos e a priorização de controles. Essa abordagem integrada reforça a robustez metodológica da proposta e está alinhada com práticas reconhecidas.

Ao integrar essas duas abordagens, foi possível agrupar os riscos de IA Generativa em quatro grupos com base na quantidade de conexões entre os riscos, conforme Tabela C1 (ver Anexo C). O referido agrupamento de padrões quantitativos reflete diferentes níveis de contribuição na mitigação dos riscos de negócio, com base na conectividade entre os riscos IA-G e os riscos de negócio. Essa divisão permite uma análise mais granular: o Grupo 1, com maior número de cruzamentos; o Grupo 2 e o Grupo 3, com níveis intermediários de conectividade; e o Grupo 4, sem cruzamentos. Essa segmentação permite que diferentes cenários sejam montados com base na efetividade de cada grupo, e pode ser aplicada a qualquer conjunto de dados que relacione riscos e controles.

Este estudo adotou como premissa que o número de interconexões entre os riscos está diretamente relacionado ao impacto no negócio, como um todo. Ou seja, considerou que quanto maior a quantidade de conexões estabelecidas, maior tende a ser a relevância estratégica e operacional da conectividade para o negócio. Essa abordagem permitiu associar a densidade de interconexões à criticidade dos processos envolvidos. Conforme a necessidade do negócio, pode-se aplicar o peso, tal como: $\text{Impacto} = (\text{Interconexões} \times w1)$, onde: $w1$ seria o peso atribuído conforme a importância relativa de cada fator.

Em relação à classificação dos grupos de conectividade, esta foi realizada com base em intervalos percentuais fixos, sem aplicação de pesos diferenciados entre os critérios considerados. Dessa forma, cada componente da conectividade contribui de forma igual para o resultado. A organização dos grupos por conectividade foi realizada com base no

percentual de riscos de negócio mitigados por cada grupo, conforme Tabela D1 (ver Anexo D). O Grupo 1, com 70% de cobertura, foi classificado como de Alta Conectividade, indicando forte impacto na prevenção e mitigação dos principais riscos. O Grupo 2, com 40%, foi classificado como de Média Conectividade, enquanto o Grupo 3, com 20%, recebeu a classificação de Baixa Conectividade. Por fim, o Grupo 4, que não contribuiu para a mitigação de riscos, foi classificado como de Sem Conectividade.

Considerando a hipótese da organização dos riscos em 4 (quatro) grupos distintos, para se estabelecer todas as possibilidades de priorização de implementação, estamos lidando com o conceito de permutação — ou seja, todas as formas possíveis de ordenar esses grupos. Portanto, existem 24 possibilidades de priorização, demonstradas nas combinações possíveis, na Tabela E1 (ver Anexo E). As combinações entre os grupos de conectividade configuram possíveis cenários estratégicos para a seleção de controles priorizados. Dentre eles, o Cenário 1 se destaca como forte candidato, uma vez que a implementação dos controles associados ao Grupo 1 responde a 70% dos riscos de negócio identificados no período T1. Já os 30% restantes, vinculados ao Grupo 2, são endereçados no período T2, resultando em uma cobertura acumulada dos riscos de negócio. Verifica-se a ocorrência de sobreposição entre os grupos, característica que contribui para ampliar a abrangência e conferir maior agilidade na resposta às vulnerabilidades identificadas.

Nesse sentido, a Tabela 5 apresenta a cobertura estimada dos riscos de negócio, por grupo de conectividade, evidenciando sobreposições. O Grupo 1 concentra a maior parte (70%), seguido pelo Grupo 2 (40%) e pelo Grupo 3 (20%), enquanto o Grupo 4 não apresenta cobertura. Como um mesmo risco pode estar associado a diferentes grupos, os percentuais expressam abrangência relativa quando considerados em conjunto.

Tabela 5 - Cobertura dos riscos de negócio por grupo de conectividade com sobreposição

Grupo	Cobertura Estimada (%)	Observação
Grupo 1	70%	Pode incluir riscos também presentes em outros grupos
Grupo 2	40%	Sobreposição parcial com Grupo 1
Grupo 3	20%	Cobertura complementar
Grupo 4	0%	Sem cobertura no cenário atual

Fonte: Elaborado pelos autores

A Tabela F1 (vide Anexo F) apresenta a aplicação prática da hipótese de adoção do Cenário 1 como estratégia de priorização de controles para mitigação de riscos de negócio relacionados à IA-G. Cada linha da tabela detalha um risco de negócio - RNEG, incluindo sua descrição, os grupos de conectividade afetados, o número de conexões envolvidas, o intervalo temporal estimado para sua ocorrência e as respectivas consequências. A análise evidencia que os riscos mais críticos, como o vazamento de informações sigilosas (RNEG 2) e a divulgação antecipada de decisões (RNEG 1), estão concentrados nos Grupos 1 e

2, que juntos oferecem cobertura significativa no período T1. A presença de múltiplas conexões e consequências associadas reforça a relevância estratégica desses riscos e justifica a priorização dos controles vinculados aos grupos mencionados. Essa abordagem permite alinhar a densidade de interconexões à criticidade dos processos, contribuindo para uma resposta mais eficiente e tempestiva às vulnerabilidades identificadas.

A adoção da hipótese de priorização com base no Cenário 1 representa uma estratégia orientada pela efetividade na mitigação dos riscos de negócio. Nesse cenário, os grupos são distribuídos conforme sua conectividade: o Grupo 1, com 70% de cobertura, ocupa a Prioridade 1, refletindo sua alta relevância e impacto direto na redução dos riscos estratégicos. O Grupo 2, com 30%, é alocado como Prioridade 2, atuando como suporte complementar. Já os Grupos 3 e 4, com baixa ou nenhuma conectividade, são posicionados nas Prioridades 3 e 4, respectivamente, indicando menor influência na mitigação. Essa estrutura hierárquica permite uma alocação direcionada de recursos e esforços, priorizando controles sobre os grupos com maior potencial de resposta aos riscos de negócio e pode ser replicada em diferentes contextos organizacionais, conforme demonstrado na Tabela 6.

Tabela 6 – Riscos Priorizados para Seleção de Controles

Seq	Risco de Negócio (RNEG's)		Riscos de IA-G (RGAI's)
1º	RNEG2	Vazamento de informações sigilosas, protegidas por segredo de Justiça ou dados pessoais	IA-G (2-4-7-8) Confabulação Privacidade de Dados Integridade da Informação Segurança da Informação
2º	RNEG1	Divulgação antecipada de votos, determinações ou decisões	IA-G (2-7-8) Confabulação Integridade da Informação Segurança da Informação
3º	RNEG3	Emissão ou alteração não autorizada de determinações ou decisões	IA-G (2-7-8) Confabulação Integridade da Informação Segurança da Informação
4º	RNEG9	Julgamentos legítimos, porém, com base em elementos adulterados	IA-G (2-7) Confabulação Integridade da Informação
5º	RNEG6	Perda de informações	IA-G (2-7-8) Confabulação Integridade da Informação Segurança da Informação
6º	RNEG10	Espionagem de outras nações e/ou grupos de interesse	IA-G (8) Segurança da Informação
7º	RNEG4	Interrupção da prestação jurisdicional	IA-G (8) Segurança da Informação

8º	RNEG8	Assuntos indesejados ou inadequados em determinações e decisões	IA-G (1-2-3-9-10-11) Informações CBRN Confabulação Recomendações Perigosas ou Violentas Propriedade Intelectual Conteúdo Obsceno, Degradante e/ou Abusivo Toxicidade, Viés e Homogeneização
9º	RNEG5	Previsibilidade ou manipulação da distribuição dos processos	IA-G (6-11-12) Configuração Humano-IA Toxicidade, Viés e Homogeneização Cadeia de Valor e Integração de Componentes
10º	RNEG7	Parcialidade ou favorecimentos pessoais	IA-G (4-6-10) Privacidade de Dados Configuração Humano-IA Conteúdo Obsceno, Degradante e/ou Abusivo

Fonte: Elaborado pelos autores

4.1 Discussão dos Resultados

A análise dos riscos de negócio (Tabela 1) e dos riscos de IA-G (Tabela 2), foi significativamente enriquecida pela aplicação integrada das técnicas *bow-tie*, conforme a Tabela A1 (vide Anexo A) e do *heat map*, conforme Figura 1 (vide Anexo B), resultando nas diretrizes para seleção de controles prioritários na mitigação de riscos associados à IA-G, apresentadas, em modelo conceitual, na Tabela 3 e na Tabela 4

A técnica *bow-tie* permitiu estruturar visualmente as relações de causa e efeito entre vulnerabilidades da IA-G e os riscos de negócio, posicionando cada risco como evento central e associando os controles recomendados pela publicação NIST.AI.600-1 (2024) como barreiras preventivas e mitigadoras.

Essa abordagem facilitou a identificação dos riscos com maior número de cruzamentos, servindo como base para a classificação dos grupos por conectividade e para a formulação de diretrizes mais precisas na priorização dos controles.

Complementarmente, o *heat map* agregou valor ao proporcionar uma visualização rápida e intuitiva das interseções entre riscos, revelando padrões e correlações que seriam menos evidentes em representações tabulares. Essa técnica permitiu destacar com clareza os pontos de maior concentração de riscos e controles, acelerando a compreensão das relações críticas e fortalecendo a fundamentação analítica para decisões estratégicas.

A análise integrada das Tabelas C1 a F1, apresentadas nos Anexos C a F, revela uma estrutura para a priorização de controles frente aos riscos de IA-G no contexto do processo de despachos e decisões no sistema judiciário brasileiro. A Tabela C1 (Anexo C) apresenta as conexões entre os riscos de IA-G e os riscos de negócio, permitindo a identificação de agrupamentos com base na intensidade das interações. Esses agrupamentos caracterizados na Tabela D1 (Anexo D), detalham os grupos formados a partir da análise de impacto, número de conexões e aplicação de pesos conforme a

relevância estratégica dos riscos para o negócio. A Tabela E1 (Anexo E) explora cenários possíveis de combinação entre os grupos, simulando trajetórias de mitigação com diferentes níveis de cobertura e eficiência. Por fim, a Tabela F1 (Anexo F) aplica a hipótese de adoção do Cenário 1, demonstrando como a ativação sequencial dos grupos pode responder de forma escalonada aos riscos de negócio, considerando o intervalo temporal e as consequências associadas. Essa abordagem permitiu alinhar conectividade, impacto e tempo de resposta, oferecendo suporte à tomada de decisão baseada em critérios técnicos, adaptáveis às necessidades organizacionais e passíveis de personalizações por meio da atribuição de pesos específicos conforme o contexto estratégico.

A aplicação do Cenário 1 demonstra uma sequência estratégica de ativação dos grupos de conectividade — iniciando pelo Grupo 1, seguido pelos Grupos 2, 3 e 4 — que permite uma cobertura progressiva dos riscos de negócio. Conforme evidenciado na Tabela E1 (vide Anexo E), os grupos com maior conectividade (Grupos 1 e 2) são responsáveis pela mitigação da maior parte dos riscos críticos nos primeiros intervalos de tempo (T1 e T2), enquanto os grupos com menor conectividade (Grupos 3 e 4) atuam de forma complementar em etapas posteriores.

Essa estrutura reforça a relevância da priorização proposta e está alinhada à lógica de resposta ágil e escalonada às vulnerabilidades identificadas. A Tabela F1 (vide Anexo F) reforça essa lógica ao detalhar a associação entre os riscos de negócio e os riscos IA-G, permitindo a seleção direcionada de controles preventivos e mitigadores com base nas recomendações do NIST.AI.600-1 (2024).

Os riscos de IA Generativa vinculados diretamente aos riscos de negócio fortalecem a fundamentação técnica para a aplicação dos controles. Essa estrutura integrada — que combina conectividade, temporalidade e impacto — oferece uma base relevante para o estabelecimento de diretrizes para seleção de controles preventivos e mitigadores no processo de despachos e decisões no sistema judiciário.

Em relação ao risco “RGAI5 - Impacto Ambiental”, embora não esteja diretamente correlacionado aos riscos de negócio do processo de despachos e decisões no sistema judiciário, o consumo de energia e recursos naturais — como eletricidade e água — durante o treinamento e a operação dos modelos, afeta a sustentabilidade da infraestrutura tecnológica (CNJ, 2024). Ou seja, além dos riscos individuais, a IA-G impõe desafios coletivos, como impactos ambientais.

Essa análise reforça a importância de estratégias de priorização, na medida em que contribui para a possibilidade objetiva de utilizar o número de conexões entre os riscos e o momento ideal de ativação de cada grupo. A estrutura temporal dos cenários permite simular diferentes trajetórias de mitigação, oferecendo suporte à tomada de decisão baseada em impacto incremental e otimização de recursos.

Além disso, a abordagem admite a aplicação de pesos diferenciados aos grupos ou riscos, conforme a importância relativa de cada fator para o negócio, permitindo adaptações dinâmicas na priorização de controles de acordo com o contexto da organização.

5. Conclusões e Trabalhos Futuros

Conclui-se que este estudo atingiu seu objetivo ao relacionar os principais fatores de risco que podem afetar o processo de despachos e decisões no sistema judiciário, com os riscos no uso de IA Generativa (NIST, 2024), tornando possível propor a seleção de controles para mitigar os referidos riscos de negócio traçando diretrizes para a implantação segura e responsável desse tipo de tecnologia nas plataformas e sistemas judiciais.

As diretrizes apresentadas visam fortalecer a governança institucional e promover maior segurança jurídica, contribuindo para decisões alinhadas aos princípios éticos e legais que regem o sistema judiciário. Os resultados evidenciaram a relevância da abordagem proposta, estruturada por meio de diretrizes, conforme demonstrado na Tabela 3 e exemplificada por meio da Tabela 4. A integração das técnicas *bow-tie*, conforme Tabela A1 (vide Anexo A) e *heat map*, conforme Figura B1 (vide Anexo B), permitiu representar visualmente as relações causais entre vulnerabilidades da IA-G e os riscos de negócio, identificar interseções relevantes, conforme Tabela C1 (vide Anexo C) e fundamentar a formação de grupos por conectividade, conforme Tabela D1 (vide Anexo D). A construção de cenários de priorização, conforme Tabela E1 (vide Anexo E) e a seleção estratégica de controles, conforme Tabela F1 (vide Anexo F) e a Tabela 6, demonstrou que a metodologia é replicável, adaptável e orientada por evidências.

A abordagem proposta eleva o padrão de governança, alinhando o uso da IA-G aos princípios de justiça, responsabilidade e confiança pública. Nesse contexto, faz-se necessária a participação ativa de magistrados, servidores e advogados, aliada à consideração de impactos ambientais e à necessidade de auditorias contínuas, reforça a importância de uma governança multidisciplinar e adaptável (Bender, Gebru & McMillan-Major, 2021; Brundage et al., 2018; CNJ, 2024).

Nessa linha, destacam-se os recentes avanços no Judiciário brasileiro, como a Resolução nº 615, de 11 de março de 2025, que estabelece diretrizes para o desenvolvimento, utilização e governança de soluções de inteligência artificial no Poder Judiciário. A resolução visa garantir a transparência, segurança e supervisão no uso de tecnologias de IA exigindo que os tribunais adotem fontes de dados seguras e auditáveis. Além disso, a resolução regulamenta a obrigatoriedade do uso de requisitos para sistemas informatizados no âmbito judiciário. Essa nova normativa é um passo importante para a modernização e a responsabilidade no uso de IA no setor público (CNJ, 2025).

Experiências internacionais reforçam essa perspectiva. Na Alemanha, sistemas como OLGA e Frauke automatizam tarefas repetitivas e reduzem o tempo de julgamento em até 50% (IBM, 2024). No Reino Unido e nos Países Baixos, a IA é usada para resolução de disputas online e previsão de resultados judiciais, enquanto a França explora a justiça preditiva para acelerar julgamentos (CEPEJ, 2024). Essas práticas mostram que o uso responsável da IA — com transparência, supervisão humana e auditorias — é essencial para garantir a imparcialidade e a confiança no sistema.

Entre as principais limitações destaca-se o contexto dos principais riscos do processo de despachos e decisões no sistema judiciário brasileiro.

Como trabalhos futuros, recomenda-se expandir a aplicação da metodologia para outros contextos organizacionais, como saúde, educação e segurança pública, testando sua escalabilidade e adaptabilidade a diferentes tipos de riscos e estruturas operacionais. Além disso, a automatização parcial do processo de agrupamento e priorização, por meio de sistemas de apoio à decisão, também representa uma oportunidade promissora para ampliar a eficiência e a aplicabilidade da proposta em ambientes complexos e dinâmicos.

Em síntese, este estudo contribui para o fortalecimento de uma cultura institucional orientada pela governança e gestão de riscos dinâmica e adaptativa, essencial diante das rápidas transformações digitais que impactam o sistema de justiça. Ao propor diretrizes para a seleção de controles prioritários para a mitigação de riscos de IA-G em despachos e decisões no sistema judiciário, busca-se assegurar a integridade das decisões judiciais, preservar direitos fundamentais e reforçar a legitimidade institucional. A articulação entre os referenciais internacionais, e os normativos nacionais, como a Resolução CNJ nº 615/2025, evidencia a importância de uma governança responsável da IA, capaz de equilibrar inovação tecnológica com segurança jurídica e transparência.

Referências

- Almada, M., & Zanatta, R. A. F. (2024). Inteligência artificial, direito e pesquisa jurídica. *Revista USP*, (141), 51–64. <https://doi.org/10.11606/issn.2316-9036.i141p51-64>
- Alvares, N. O. (2012). A informatização do processo judicial e o acesso à justiça. (Monografia). Faculdade de Ciências Jurídicas e de Ciências Sociais, UniCEUB, Brasília
- Alves, L. F., Georg, A. M., & Nunes, J. C. (2023). Judiciário sob ataque hacker: Riscos de negócio para segurança cibernética em tribunais brasileiros. *Revista Ibérica de Sistemas e Tecnologias de Informação*, pp. 344-357
- Alves, R. S., Barbacena da Silva, J. P., Ribeiro Júnior, L. A., & Nunes, R. R. (2025). Enhancing cybersecurity in the judiciary: Integrating additional controls into the CIS framework. *Computers & Security*, 157, Article 104584. <https://doi.org/10.1016/j.cose.2025.104584>
- Alves, R. S., da Silva, J. P. B., Ribeiro Junior, L. A., & Nunes, R. R. (2025). Enhancing cybersecurity in the judiciary: Integrating additional controls into the CIS framework. *Computers & Security*, 157, 104584. <https://doi.org/10.1016/j.cose.2025.104584>

Alves, R. S., Georg, M. A., & Nunes, R. R. (2023). Judiciário sob ataque hacker: Riscos de negócio para segurança cibernética em tribunais brasileiros. *Revista Ibérica de Sistemas e Tecnologias de Informação*, 2023(1), 344-357

Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., Garcia, S., Gil-Lopez, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 136-157. <https://arxiv.org/abs/1910.09337>

Associação Brasileira de Normas Técnicas. (2023). Tecnologia da informação - Inteligência artificial - Conceitos e terminologia (ABNT NBR ISO/IEC 22989:2023).

Associação Brasileira de Normas Técnicas. (2024). Tecnologia da Informação – Inteligência Artificial – Sistema de Gestão (ABNT NBR ISO/IEC 42001:2024).

Associação Brasileira de Normas Técnicas. Segurança da Informação, segurança cibernética e proteção à privacidade – Requisitos (ABNT NBR ISO IEC 27001).

Associação Brasileira de Normas Técnicas. Segurança da Informação, segurança cibernética e proteção à privacidade – Orientações para gestão de riscos de segurança da informação (ABNT NBR ISO IEC 27005).

Associação Brasileira de Normas Técnicas. Gestão de Riscos – Diretrizes (ABNT NBR ISO IEC 31000).

Associação dos Magistrados Brasileiros (AMB). (2007). Judiciário brasileiro em perspectiva: análise da Associação dos Magistrados Brasileiros baseada em relatórios do Supremo Tribunal Federal, do Conselho Nacional de Justiça e do Banco Mundial <http://www.amb.com.br/>

Bender, E. M., Gebru, T., & McMillan-Major, A. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency* (pp. 610-623)

Brasil. Decreto nº 9.203, de 22 de novembro de 2017. Dispõe sobre a política de governança da administração pública federal direta, autárquica e fundacional. *Diário Oficial da União*.

Brasil. (2021). Estratégia Brasileira de Inteligência Artificial (EBIA). Ministério da Ciência, Tecnologia e Inovações. <https://www.gov.br/mcti/pt-br/assuntos/noticias/ebia-lanca-estrategia-brasileira-de-inteligencia-artificial>

Brundage, M., Avin, S., Wang, J., Belfield, H., Krueger, G., Hadfield, G., & Legassick, S. (2018). The malicious use of artificial intelligence: Forecasting, prevention, and mitigation

CEPEJ. Comissão Europeia para a Eficiência da Justiça. (2024). Acesso em 08 de setembro de 2024. <https://www.coe.int/en/web/cepej/practical-examples-of-ai-implemented-in-other-countries>

Chesney, R., & Citron, D. (2019). Deepfakes and the New Disinformation War: The Coming Age of Post-Truth Geopolitics. *Foreign Affairs*. <https://www.foreignaffairs.com/articles/2019-08-12/deepfakes-and-new-disinformation-war>

Conselho Nacional de Justiça. (2021). Justiça 4.0. Acesso em 21 de maio de 2022. <https://www.cnj.jus.br/tecnologia-da-informacao-e-comunicacao/justica-4-0/>

Conselho Nacional de Justiça. (2024). Especialistas discutem segurança cibernética da Justiça no Link CNJ*. Acesso em 07 de setembro de 2024. <https://www.cnj.jus.br/especialistas-discutem-protecao-cibernetica-da-justica-no-link-cnj/>

Conselho Nacional de Justiça. (2024). Estatísticas do Poder Judiciário.

Conselho Nacional de Justiça. (2024). Justiça 4.0: Inteligência Artificial está presente na maioria dos tribunais brasileiros. Acesso em 08 de setembro de 2024. <https://www.cnj.jus.br/justica-4-0-inteligencia-artificial-esta-presente-na-maioria-dos-tribunais-brasileiros/>

Conselho Nacional de Justiça. (2024). Programa Justiça 4.0 divulga resultados de pesquisa sobre IA no Judiciário brasileiro. Acesso em 08 de setembro de 2024. <https://www.cnj.jus.br/programa-justica-4-0-divulga-resultados-de-pesquisa-sobre-ia-no-judiciario-brasileiro/>

Conselho Nacional de Justiça. (2024). Relatório de Pesquisa sobre o Uso da Inteligência Artificial Generativa no Poder Judiciário Brasileiro. Acesso em 20 de junho de 2025. [cnj-relatorio-de-pesquisa-ia-g-pj.pdf](#)

Crawford, K., & Calo, R. (2016). There is a Blind Spot in AI Research. *Nature*, 538(7624), 311-313

Duffourc, M., & Gerke, S. (2023). Generative AI in health care and liability risks for physicians and safety concerns for patients. *JAMA*, 330(14), 1327-1328.

European Union Agency for Cybersecurity (ENISA). (2023). AI Cybersecurity: Towards Trustworthy AI Development. <https://www.enisa.europa.eu/publications/ai-cybersecurity-towards-trustworthy-ai-development>

Gartner. (2024). O uso indevido da IA generativa diminui o valor da IA nas organizações. Acesso em 08 de setembro de 2024. <https://www.gartner.com.br/pt-br/artigos/quando-nao-usar-a-ia-generativa>

Greshake, K., et al. (2023). More than you've asked for: A Comprehensive Analysis of Indirect Prompt Injection Attacks in LLM Applications. <https://arxiv.org/abs/2302.05737>

Hino, M. C., & Cunha, M. A. (2020). Adoção de tecnologias na perspectiva de profissionais de direito. *Revista Direito GV*, 16(1), e1952. <https://doi.org/10.1590/2317-6172201952>

Hoch, P. A., & Engelmann, W. (2023). Regulação da inteligência artificial no judiciário brasileiro e europeu. <https://ojs.unifor.br/rpen/article/view/14263>

IBM. (2024). Judicial systems are turning to AI to help manage vast quantities of data and expedite case resolution. Acesso em 08 de setembro de 2024. <https://www.ibm.com/blog/judicial-systems-are-turning-to-ai-to-help-manage-its-vast-quantities-of-data-and-expedite-case-resolution/>

Instituto Brasileiro de Governança Corporativa. (2020). Governança corporativa: princípios e melhores práticas. IBGC.

International Organization for Standardization. (2019). ISO 31010:2019 — Risk management — Risk assessment techniques. ISO. <https://www.iso.org/standard/72140.html>

Martins, T. d. (2021, 20 de janeiro). Acesso à Justiça e pandemia. Revista Jus Navegandi, 6412. Acesso em 21 de maio de 2022. <https://jus.com.br/artigos/88048/acesso-a-justica-e-pandemia>

Moreira, F. R., Da Silva Filho, D. A., Nze, G. D., De Sousa Junior, R. T., & Nunes, R. R. (2021). Evaluating the Performance of NIST's Framework Cybersecurity Controls Through a Constructivist Multicriteria Methodology. IEEE Access, 9, 129605-129618. <https://doi.org/10.1109/access.2021.3113178>

National Institute of Standards and Technology (NIST). (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>

National Institute of Standards and Technology. (2024). Artificial intelligence risk management framework: Generative artificial intelligence profile (NIST AI-600-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI-600-1>

Oliveira, C. G. B. (2022). Desafios da regulação do digital e da inteligência artificial no Brasil. Revista USP, (135), 31–44. <https://revistas.usp.br/revusp/article/view/206257>

Pinto, M. V. (2025). Trilha de Implantação de Governança de Dados para IA (TI Exames, 2025).

Russel, S., & Norvig, P. (2010). Artificial Intelligence: a modern approach. Pearson Education

Shokri, R., Stronati, M., Song, C., & Shmatikov, V. (2017). Membership Inference Attacks Against Machine Learning Models. In Proceedings of the 2017 IEEE Symposium on Security and Privacy (pp. 3-17). <https://arxiv.org/abs/1610.05820>

Stanford Encyclopedia of Philosophy. (2020). Ethics of Artificial Intelligence. <https://plato.stanford.edu/entries/ethics-ai/>

STF. (2023). Acessado em 04 de setembro de 2024. <https://noticias.stf.jus.br/postsnoticias/stf-conclui-chamamento-publico-para-uso-de-inteligencia-artificial/>

Tredinnick, L., & Laybats, C. (2023). The dangers of generative artificial intelligence. Business Information Review, 40(2), 46-48

UNESCO. (2024). AI and the Rule of Law: Capacity Building for Judicial Systems. Acesso em 08 de setembro de 2024. <https://www.unesco.org/en/artificial-intelligence/rule-law/mooc-judges>

Wach, K., Duong, C. D., Ejdy, J., Kazlauskaitė, R., Korzynski, P., Mazurek, G., ... & Ziemba, E. (2023). The dark side of generative artificial intelligence: A critical analysis of controversies and risks of ChatGPT. Entrepreneurial Business and Economics Review, 11(2), 7-30

Wallace, E., et al. (2021). Universal Adversarial Triggers for Text Classification. <https://arxiv.org/abs/1908.07125>

Wilkinson, L., & Friendly, M. (2009). The history of the cluster heat map. *The American Statistician*, 63(2), 179–184. <https://doi.org/10.1198/tas.2009.0033>

ANEXO A

Tabela A1 - Relação entre Riscos de negócio, Riscos de IA-G, Controles e Funções NIST com Aplicação da Técnica bow-tie

N	Causas	Barreiras Preventivas	Evento Central	Consequências	Barreiras Mitigadoras
	Riscos de IA-G	Exemplos de Controles	Riscos de Negócio	Riscos de Negócio	Exemplos de Controles
1	Confabulação RGAI2 Integridade da Informação RGAI7 Segurança da Informação RGAI8	Revisão e verificação de fontes (MS) Validação de informações e fontes (MS, MG) Salvaguardas de integridade, controle de acesso, criptografia (MP, MS) Monitoramento de conteúdo sensível (MG, MS) Segurança cibernética e controle de vulnerabilidades (MS)	Divulgação antecipada de votos, determinações ou decisões RNEG1	Vazamento de informações sigilosas RNEG2 Perda de informações RNEG6 Emissão não autorizada RNEG3	Auditoria e monitoramento (MG, MS) Resposta a incidentes (GV, MG) Revisão contínua de logs e <i>outputs</i> (MG)
2	Confabulação RGAI2 Privacidade de Dados RGAI4 Integridade da Informação RGAI7 Segurança da Informação RGAI8	Proteção de dados pessoais (MP, MS) Técnicas de anonimização, controles de acesso (MP, MS) Segmentação de rede (MP) Monitoração contínua Salvaguarda da integridade (MS)	Vazamento de informações sigilosas, protegidas por segredo de Justiça ou dados pessoais RNEG2	Parcialidade ou favorecimento RNEG7 Espionagem RNEG10	Notificação e rastreamento de incidentes (MG) Gestão de crise e contenção de vazamentos (MG, MP) Comunicação à autoridade (GV)
3	Confabulação RGAI2 Integridade da Informação RGAI7 Segurança da Informação RGAI8	Assinatura digital e controle de versões (MS) Validação cruzada IA-humano (MS, MG) Segregação de funções (GV) Monitoração de integridade (MS)	Emissão ou alteração não autorizada de determinações ou decisões RNEG3	Decisões com base em elementos adulterados RNEG9 Divulgação antecipada RNEG1	Auditoria forense (MG) Revisão e correção de registro (MG, MS) Revisão processual (MG)

4	Segurança da Informação RGAI8	Firewalls (MS, MG), backup em tempo real (MS, MG), plano de continuidade (GV, MS), redundância de infraestrutura (MS, MG)	Interrupção da prestação jurisdicional RNEG4	Perda de informações RNEG6 Assuntos inadequados RNEG8	Restauração rápida (MS) Redundância amplificada (MP) Execução de plano de contingência (MG)
5	Configuração Humano-IA RGAI6 Toxicidade, Viés e Homogeneização RGAI11 Cadeia de Valor e Integração de Componentes RGAI12	Transparência de algoritmos, interfaces transparentes (GV) Curadoria e revisão independente (MS, MG) Avaliação de fornecedores e transparência da cadeia (GV) Métricas de diversidade (MG)	Previsibilidade ou manipulação da distribuição dos processos RNEG5	Parcialidade RNEG7 Emissão não autorizada RNEG3	Inspeção externa (MG) Reversão de decisões suspeitas (MG) Contestação processual (GV)
6	Confabulação RGAI2 Integridade da Informação RGAI7 Segurança da Informação RGAI8	Backup automático e redundância (MS, MP) Verificação e controle de integridade de dados (MS, MG) Monitoramento de falhas operacionais (MG)	Perda de informações RNEG6	Interrupção jurisdicional RNEG4 Vazamento de informações RNEG2	Recuperação de backup (MG) <i>Disaster recovery</i> (MS, MG) Reconstituição parcial de dados (MG)
7	Privacidade de Dados RGAI4 Configuração Humano-IA RGAI6 Conteúdo Obsceno, Degradante e/ou Abusivo RGAI10 Toxicidade, Viés e Homogeneização RGAI11	Curadoria e validação de dados (MS, MG) Monitoramento e detecção de viés (MG) Moderação e filtros automáticos (MS) Revisão humana (GV, MS)	Parcialidade ou favorecimentos pessoais RNEG-7	Assuntos inadequados RNEG8 Emissão não autorizada RNEG3	Atualização de <i>datasets</i> (MS, MG) Auditoria e contestação processual (GV, MG) Revisão manual de outputs (MG)

8	Informações CBRN RGAI1 Recomendações Perigosas ou Violentas RGAI3 Confabulação RGAI2 Propriedade Intelectual RGAI9 Conteúdo Obsceno, Degradante e/ou Abusivo RGAI10 Toxicidade, Viés e Homogeneização RGAI11	Filtros e detecção de conteúdo crítico (MS, MG) Blacklists/whitelists temáticos (MS, MG) Moderação automática (MS) Detecção de plágio, direitos autorais (MS, GV)	Assuntos indesejados ou inadequados em determinações e decisões RNEG8	Divulgação antecipada RNEG1 Parcialidade RNEG7	Retirada de conteúdo (MG) Revisão e responsabilização disciplinar (GV, MG) Revisão manual (MG)
9	Confabulação RGAI2 Integridade da Informação RGAI7	Controle de versões, registro imutável (MS, MG) Validação cruzada IA-humano (MS) Documentação de prevalência de erros e histórico de alterações (MG)	Julgamentos legítimos, porém, com base em elementos adulterados RNEG9	Emissão/alteração não autorizada RNEG3 Perda de informações RNEG6	Investigação forense de registros (MG) Reprocessamento e auditoria pós-processo (MG, MS)
10	Segurança da Informação RGAI-8	Monitoramento de rede e de acesso (MS, MG) Segmentação de sistemas e autenticação forte (MP, MS) Treinamento de pessoal em cibersegurança (GV, MS) Detecção de anomalias (MG)	Espionagem de outras nações e/ou grupos de interesse (RNEG10)	Vazamento de informações RNEG2 Perda de informações RNEG6	Isolamento de sistemas afetados (MG) Revogação de credenciais/acesso (MS, MG) Resposta a incidentes cibernéticos (MG)

Fonte: Elaborado pelos autores

ANEXO B

Figura B1 - Mapa de Calor (Riscos de Negócio versus Riscos IA-G)

Riscos de Negócio	RGAI-1	RGAI-2	RGAI-3	RGAI-4	RGAI-5	RGAI-6	RGAI-7	RGAI-8	RGAI-9	RGAI-10	RGAI-11	RGAI-12	Total
RNEG-1 Divulgação antecipada de votos, de	0	1	0	0	0	0	1	1	0	0	0	0	3
RNEG-2 Vazamento de informações sigilosa	0	1	0	1	0	0	1	1	0	0	0	0	4
RNEG-3 Emissão ou alteração não autoriza	0	1	0	0	0	0	1	1	0	0	0	0	3
RNEG-4 Interrupção da prestação jurisdicio	0	0	0	0	0	0	0	1	0	0	0	0	1
RNEG-5 Previsibilidade ou manipulação da	0	0	0	0	0	1	0	0	0	0	1	1	3
RNEG-6 Perda de informações	0	1	0	0	0	0	1	1	0	0	0	0	3
RNEG-7 Parcialidade ou favorecimentos pe	0	0	0	1	0	1	0	0	0	1	0	0	3
RNEG-8 Assuntos indesejados ou inadequa	1	1	1	0	0	0	0	0	1	1	1	0	6
RNEG-9 Julgamentos legítimos, porém, con	0	1	0	0	0	0	1	0	0	0	0	0	2
RNEG-10 Espionagem de outras nações e ón	0	0	0	0	0	0	0	1	0	0	0	0	1

Fonte: Elaborado pelos autores

ANEXO C

Tabela C1 – Conexões e Grupos

Riscos de IA-G RGAI's (NIST, 2024)		Conexões RGAI vs RNEG's	Impacto Peso	Agrupamento
1	Informações CBRN	1 conexão	1	Grupo 3
2	Confabulação	6 conexões	1	Grupo 1
3	Recomendações Perigosas ou Violentas	1 conexão	1	Grupo 3
4	Privacidade de Dados	2 conexões	1	Grupo 2
5	Impacto Ambiental	0 conexões	1	Grupo 4
6	Configuração Humano-IA	2 conexões	1	Grupo 2
7	Integridade da Informação	5 conexões	1	Grupo 1
8	Segurança da Informação	6 conexões	1	Grupo 1
9	Propriedade Intelectual	1 conexão	1	Grupo 3
10	Conteúdo Obsceno, Degradante e/ou Abusivo	2 conexões	1	Grupo 2
11	Toxicidade, Viés e Homogeneização	2 conexões	1	Grupo 2
12	Cadeia de Valor e Integração de Componentes	1 conexão	1	Grupo 3

Fonte: Elaborado pelos autores

ANEXO D

Tabela D1 – Características dos Grupos

Grupos	Riscos IA-G	Riscos de Negócio	Análise	%Riscos Mitigados	Características
Grupo 1	RGAI2 RGAI7 RGAI8	RNEG1 RNEG2 RNEG3 RNEG4 RNEG6 RNEG9 RNEG10	A aplicação dos controles preventivos e mitigadores aos riscos IA-G (2-7-8), responde a 7 (sete) do universo dos 10 (dez) principais riscos de negócio, ou seja: 70% dos riscos de negócio estariam mitigados.	70%	Alta Conectividade
Grupo 2	RGAI4 RGAI6 RGAI10 RGAI11	RNEG2 RNEG5 RNEG7 RNEG8	A aplicação dos controles preventivos e mitigadores aos riscos IA-G (4-6-10-11), responde a 4 (quatro) do universo dos 10 (dez) principais riscos de negócio, ou seja: 40% dos riscos de negócio estariam mitigados, 30% considerando RNEG2 já contemplado no grupo 1.	40%	Média Conectividade
Grupo 3	RGAI1 RGAI3 RGAI9 RGAI12	RNEG5 RNEG8	A aplicação dos controles preventivos e mitigadores aos riscos IA-G (1-3-9-12), responde a 2 (dois) do universo dos 10 (dez) principais riscos de negócio, ou seja: 20% dos riscos de negócio estariam mitigados, 0% considerando RNEG5 e RNEG8 já contemplado no grupo 2.	20%	Baixa Conectividade
Grupo 4	RGAI5			0%	Sem Conectividade

Fonte: Elaborado pelos autores

ANEXO E

Tabela E1 – Cenários possíveis de combinações para a seleção de controles priorizados

Cenários	Prioridade 1 Intervalo Tempo 1 T1	Prioridade 2 Intervalo Tempo 2 T2	Prioridade 3 Intervalo Tempo 3 T3	Prioridade 4 Intervalo Tempo 4 T4	Mitigação Riscos Negócio Análise Tempo
Cenário 1	Grupo 1	Grupo 2	Grupo 3	Grupo 4	Alcança 70% de cobertura em T1, 100% em T2 e o grupo 3 (com baixa conectividade) já fica atendido em T2.
Cenário 2	Grupo 1	Grupo 2	Grupo 4	Grupo 3	Alcança 70% de cobertura em T1, 100% em T2 e o grupo 3 somente ficaria atendimento em T4
Cenário 3	Grupo 1	Grupo 3	Grupo 2	Grupo 4	Alcança 70% de cobertura em T1 e 90% em T2
Cenário 4	Grupo 1	Grupo 3	Grupo 4	Grupo 2	
Cenário 5	Grupo 1	Grupo 4	Grupo 2	Grupo 3	Alcança 70% de cobertura em T1 e 0 em T2
Cenário 6	Grupo 1	Grupo 4	Grupo 3	Grupo 2	
Cenário 7	Grupo 2	Grupo 1	Grupo 3	Grupo 4	Alcança 40% de cobertura em T1
Cenário 8	Grupo 2	Grupo 1	Grupo 4	Grupo 3	
Cenário 9	Grupo 2	Grupo 3	Grupo 1	Grupo 4	
Cenário 10	Grupo 2	Grupo 3	Grupo 4	Grupo 1	
Cenário 11	Grupo 2	Grupo 4	Grupo 1	Grupo 3	
Cenário 12	Grupo 2	Grupo 4	Grupo 3	Grupo 1	
Cenário 13	Grupo 3	Grupo 1	Grupo 2	Grupo 4	
Cenário 14	Grupo 3	Grupo 1	Grupo 4	Grupo 2	
Cenário 15	Grupo 3	Grupo 2	Grupo 1	Grupo 4	Alcança 20% de cobertura em T1
Cenário 16	Grupo 3	Grupo 2	Grupo 4	Grupo 1	

Cenário 17	Grupo 3	Grupo 4	Grupo 1	Grupo 2	Alcança 0% de cobertura em T1
Cenário 18	Grupo 3	Grupo 4	Grupo 2	Grupo 1	
Cenário 19	Grupo 4	Grupo 1	Grupo 2	Grupo 3	
Cenário 20	Grupo 4	Grupo 1	Grupo 3	Grupo 2	
Cenário 21	Grupo 4	Grupo 2	Grupo 1	Grupo 3	
Cenário 22	Grupo 4	Grupo 2	Grupo 3	Grupo 1	
Cenário 23	Grupo 4	Grupo 3	Grupo 1	Grupo 2	
Cenário 24	Grupo 4	Grupo 3	Grupo 2	Grupo 1	

Fonte: Elaborado pelos autores

ANEXO F

Tabela F1 – Aplicação da hipótese de adoção do cenário 1 como priorização

Seq	Riscos de Negócio – RNEG's		Grupos	Conexões	Intervalo de Tempo	Conseq.
1	RNEG2	Vazamento de informações sigilosas, protegidas por segredo de Justiça ou dados pessoais Controles dos riscos IA-G (2-4-7-8)	1 e 2	4	T1 T2	RNEG7 RNEG10
2	RNEG1	Divulgação antecipada de votos, determinações ou decisões Controles dos riscos IA-G (2-7-8)	1	3	T1	RNEG2 RNEG3 RNEG6
3	RNEG3	Emissão ou alteração não autorizada de determinações ou decisões Controles dos riscos IA-G (2-7-8)	1	3	T1	RNEG1 RNEG9
4	RNEG9	Julgamentos legítimos, porém, com base em elementos adulterados Controles dos riscos IA-G (2-7)	1	3	T1	RNEG3 RNEG6
5	RNEG6	Perda de informações Controles dos riscos IA-G (2-7-8)	1	3	T1	RNEG2 RNEG4
6	RNEG10	Espionagem de outras nações e/ou grupos de interesse Controles dos riscos IA-G (8)	1	1	T1	RNEG2 RNEG6
7	RNEG4	Interrupção da prestação jurisdicional Controles dos riscos IA-G (8)	1	1	T1	RNEG6 RNEG8
8	RNEG8	Assuntos indesejados ou inadequados em determinações e decisões Controles dos riscos IA-G (1-2-3-9-10-11)	2 e 3	6	T2 T3	RNEG1 RNEG7
9	RNEG5	Previsibilidade ou manipulação da distribuição dos processos Controles dos riscos IA-G (6-11-12)	2 e 3	3	T2 T3	RNEG3 RNEG7
10	RNEG7	Parcialidade ou favorecimentos pessoais Controles dos riscos IA-G (4-6-10)	2	3	T2 T3	RNEG3 RNEG8

Fonte: Elaborado pelos autores